

MuS: Multilingual Synergy with Shared Representations for Visual Speech Recognition

Yuheng Fan^{1,2}

Shuang Yang^{1,2}

Shiguang Shan^{1,2}

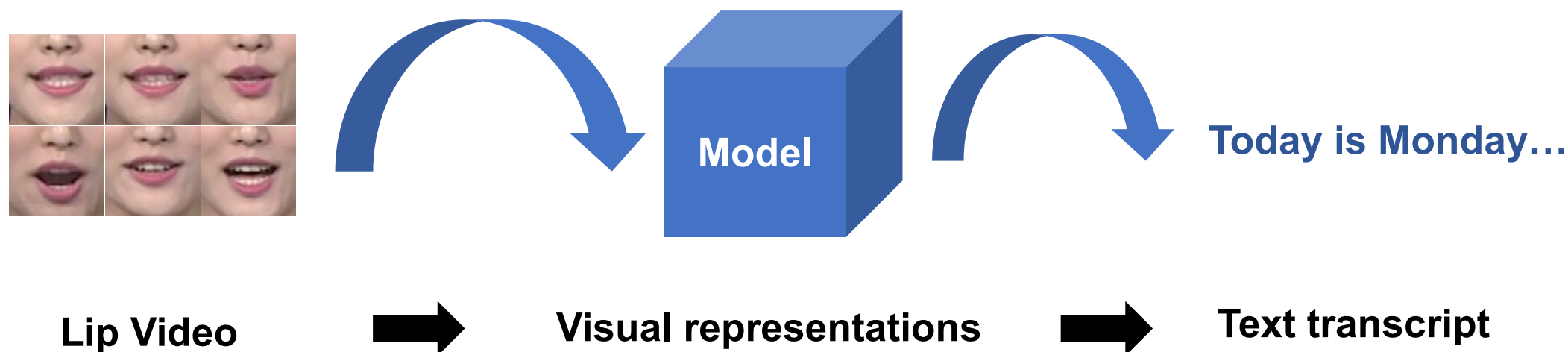
Xilin Chen^{1,2}

¹Institute of Computing Technology, Chinese Academy of Sciences

²University of Chinese Academy of Sciences



Visual Speech Recognition(VSR): Reading Speech from Lip Movements



- Speech recognition in noisy environments
- Privacy-sensitive and silent-interaction scenarios

□ From Monolingual to Multilingual VSR

Why Multilingual?



Limited per-language

Video data are limited for any single language



Scalability

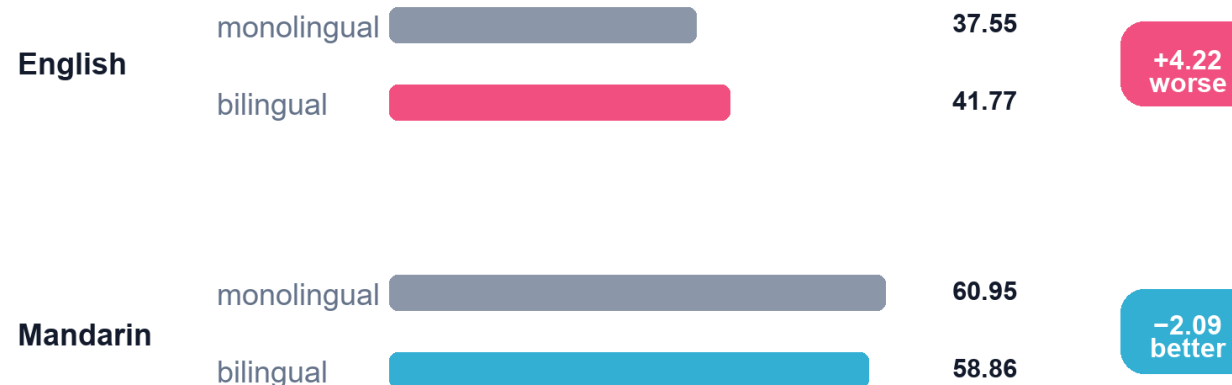
One unified model for multiple languages



Shared articulation

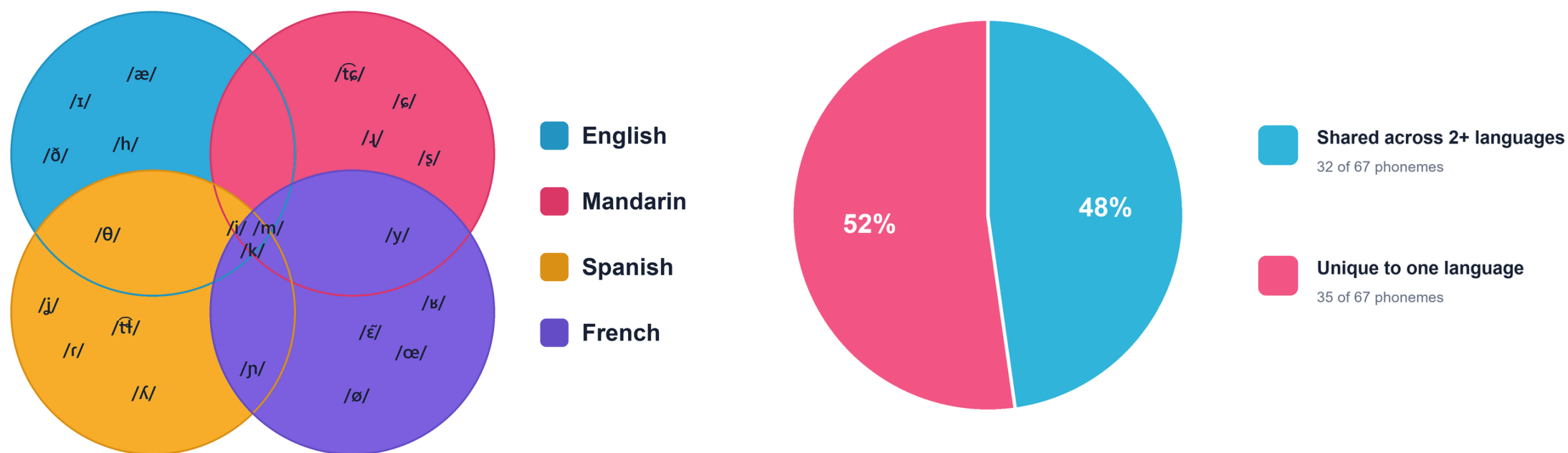
Different languages may share phonemes and corresponding lip-movement patterns

Joint training ≠ joint gains



The “curse of multilinguality”

Shared Phonetic Patterns



Cross-Lingual Phoneme Overlap

□ Shared Linguistic Structures

Different languages exhibit similar word-order and sentence patterns.

eg.



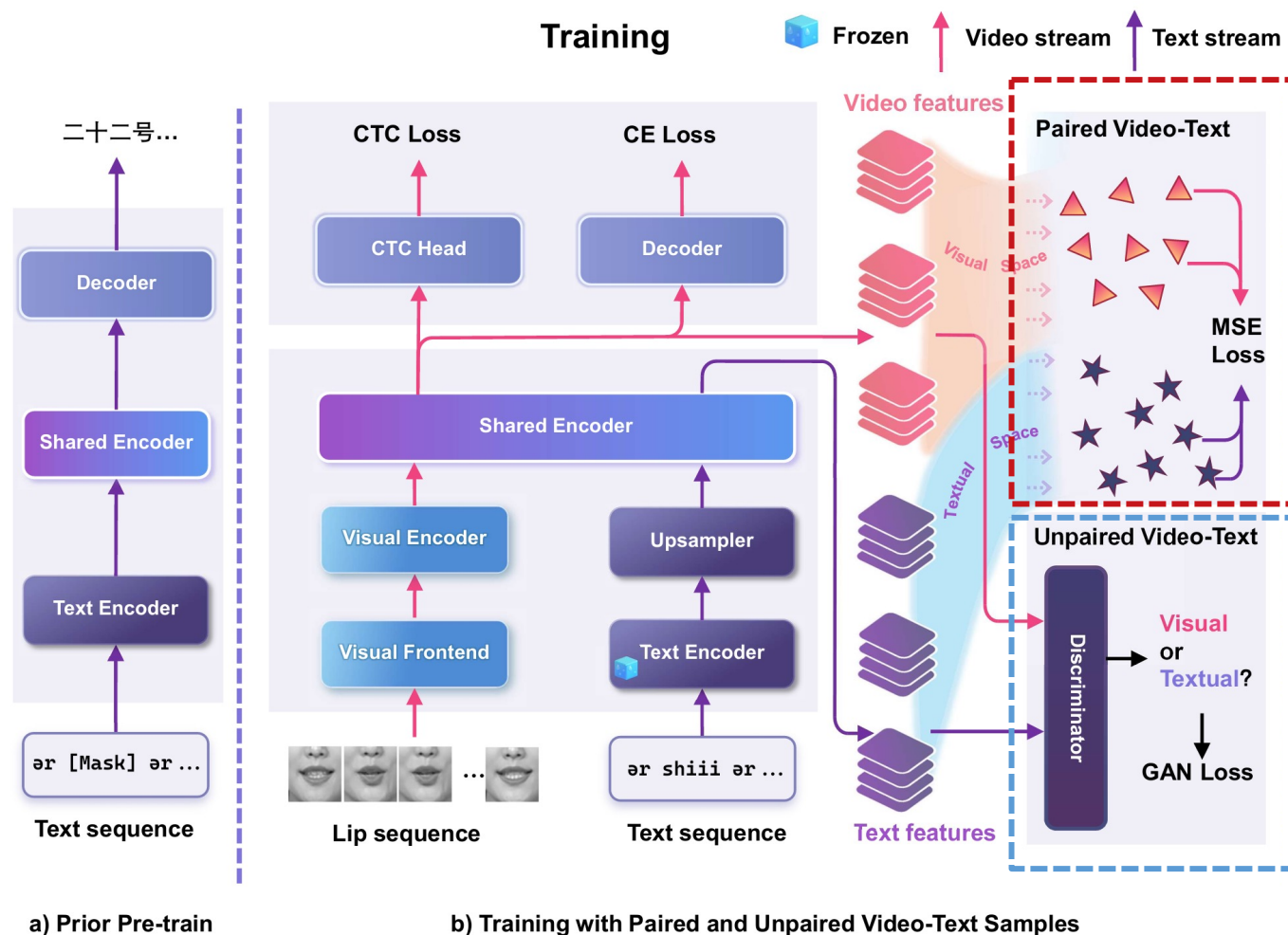
I / like / this book.

我 / 喜欢 / 这本书

Subject-Verb-Object

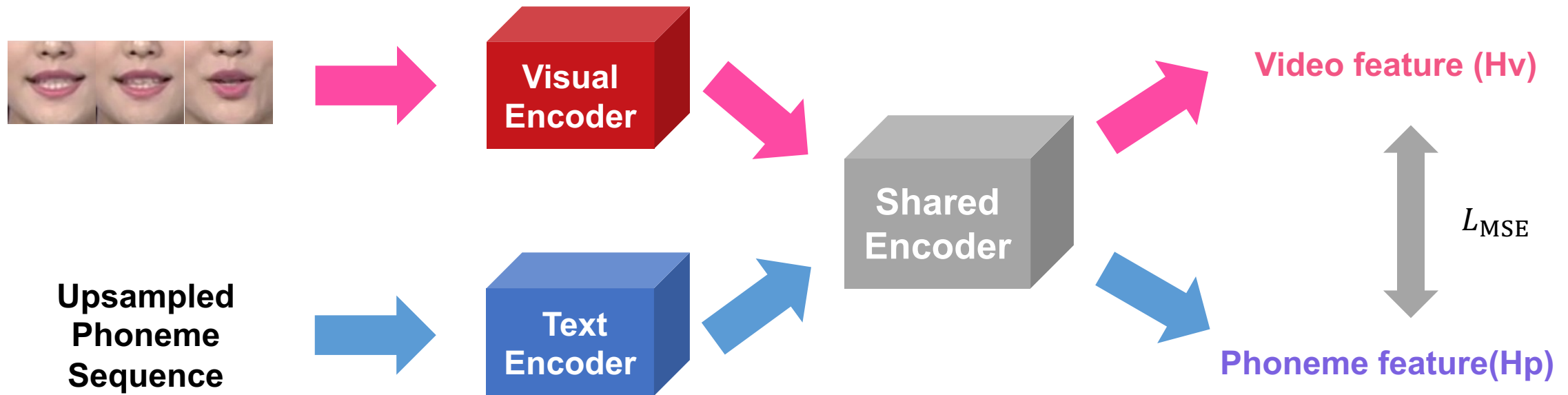
Contextual and syntactic regularities can provide transferable linguistic cues.

Overall Framework



- a) Linguistic prior pre-training
- b) Paired & Unpaired Learning

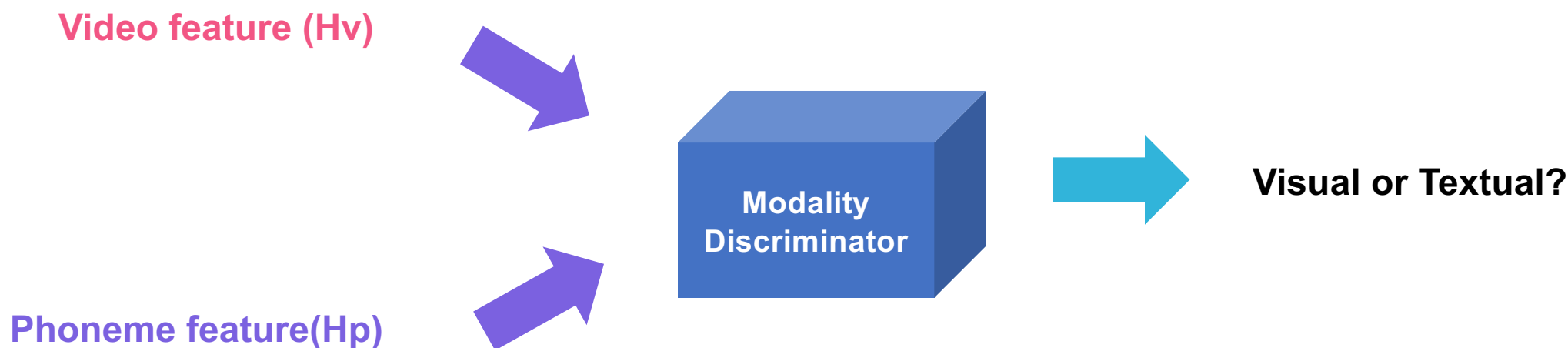
Paired Samples: MSE Loss



$$L_{MSE} = \frac{1}{|B|} \sum_{b=1}^{|B|} \|H_v - H_p\|_2^2$$

Same utterance → direct visual–phoneme representation matching

□ Unpaired Samples: Discriminate Loss



$$L_u = \min_M \max_C E_{H_p} [\log D(H_p)] + E_{H_v} [\log(1 - D(H_v))]$$

No content correspondence → align distributions, not individual instances

□ Consistent Performance Improvements Across All Languages

Table 1. Bilingual Experimental Results

Bilingual

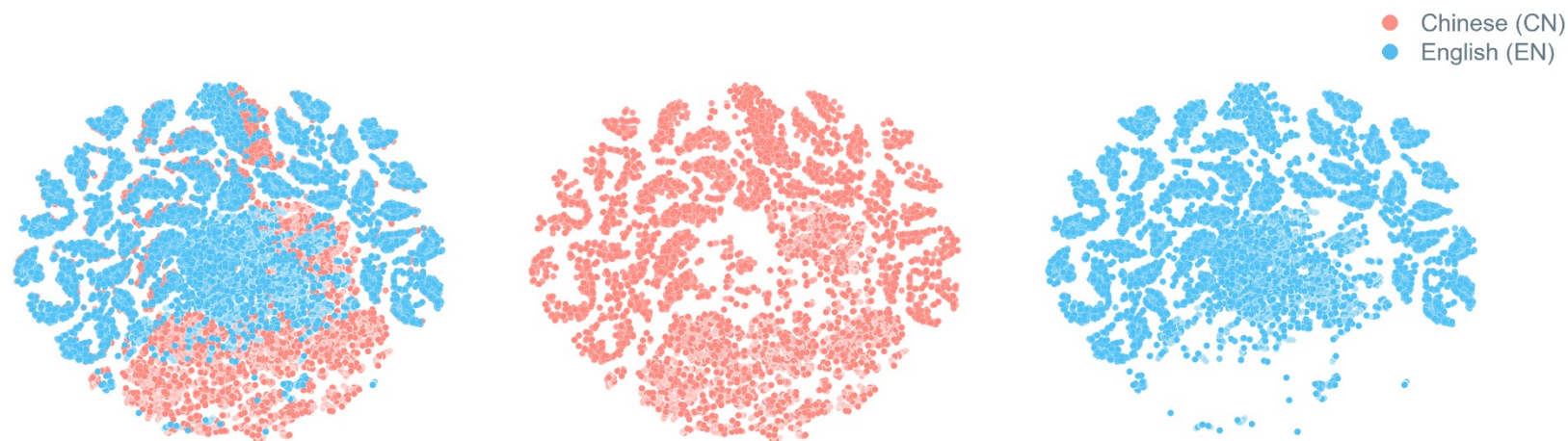
Method	Modality	WER (%)	
		en	zh
<i>Monolingual</i>			
BaseModel (en)	V	37.55	—
BaseModel (zh)	V	—	60.95
<i>Bilingual</i>			
BaseModel	V	41.77	58.86
MuS (Ours)	V + T	31.47	47.91

Table 4. Evaluation of Multilingual Visual Speech Recognition Methods.

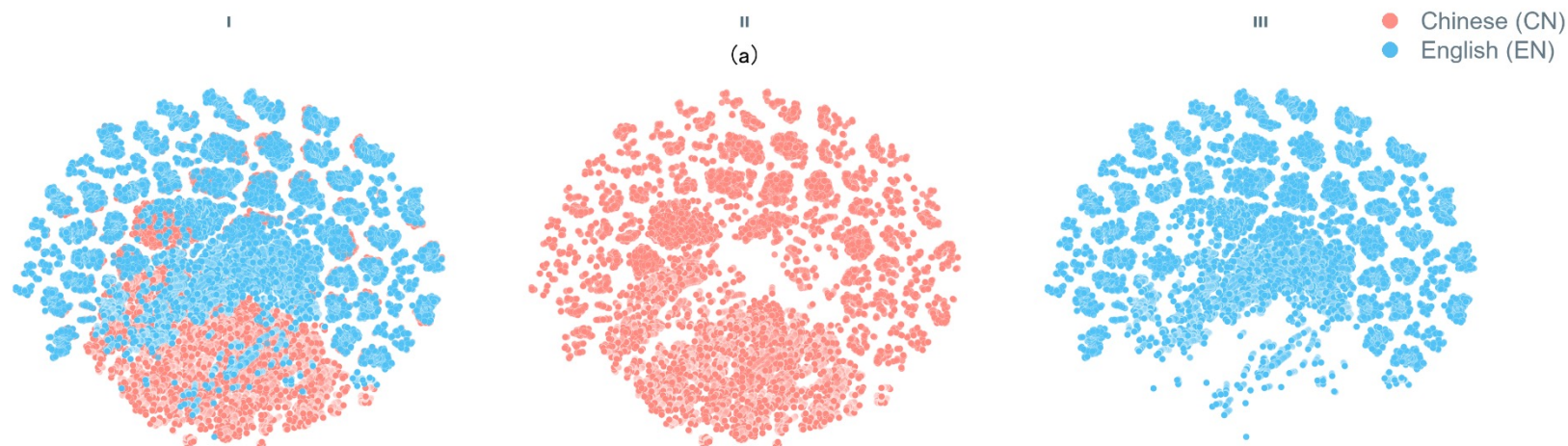
Multilingual

Method	Pre-train	Training	Single	WER (%)			
	Data	Data	Model	en	zh	es	fr
BaseModel	–	542h	✓	40.04	56.93	76.93	83.26
MuS (Ours)	–	542h	✓	31.44	47.20	65.89	75.88

Visualizing Learned Video Representations



Base model
Looser, more language-specific clusters



MuS (Ours)
More compact and language-invariant

Thanks for listening!

fanyuheng25@mailsucas.ac.cn

